



Models and Mechanisms for Spatial Data Fairness

Sina Shaham

Viterbi School of Engineering
University of Southern California
Los Angeles, California, USA
sshaham@usc.edu

Gabriel Ghinita

College of Science and Engineering
Hamad Bin Khalifa University
Qatar Foundation, Doha, Qatar
gghinita@hbku.edu.qa

Cyrus Shahabi

Viterbi School of Engineering
University of Southern California
Los Angeles, California, USA
shahabi@usc.edu

ABSTRACT

Fairness in data-driven decision-making studies scenarios where individuals from certain population segments may be unfairly treated when being considered for loan or job applications, access to public resources, or other types of services. In location-based applications, decisions are based on individual whereabouts, which often correlate with sensitive attributes such as race, income, and education.

While fairness has received significant attention recently, e.g., in machine learning, there is little focus on achieving fairness when dealing with location data. Due to their characteristics and specific type of processing algorithms, location data pose important fairness challenges. We introduce the concept of *spatial data fairness* to address the specific challenges of location data and spatial queries. We devise a novel building block to achieve fairness in the form of *fair polynomials*. Next, we propose two mechanisms based on fair polynomials that achieve individual spatial fairness, corresponding to two common location-based decision-making types: *distance-based* and *zone-based*. Extensive experimental results on real data show that the proposed mechanisms achieve spatial fairness without sacrificing utility.

PVLDB Reference Format:

Sina Shaham, Gabriel Ghinita, and Cyrus Shahabi. Models and Mechanisms for Spatial Data Fairness. PVLDB, 16(2): 167 - 179, 2022.
doi:10.14778/3565816.3565820

PVLDB Artifact Availability:

The source code, data, and/or other artifacts have been made available at <https://github.com/SinaShaham/c-Fair-Polynomials/tree/main>

1 INTRODUCTION

In the past decade, location data became an integral part of many applications, e.g., mobile apps, smart health, smart cities. Individual locations are often used in creating user profiles, or as an input for various decision-making processes (e.g., machine learning), which may affect an individual's access to public resources, loans, etc.

It is already well-understood that location bias has significant effects on underprivileged communities, e.g., in the context of transportation and housing [25]. Some public authorities designated entire geographical regions as economically-disadvantaged areas (EDA) [2]. However, while fairness has been studied recently in ML settings [6] for generic data types [16], no specific solution studies

fairness in spatial data processing. Addressing spatial data fairness presents two specific challenges:

(1) Location data may lead, intentionally or inadvertently, to exercise bias against individuals from disadvantaged backgrounds in a *stealth* fashion. While it is illegal to use race or ethnicity in a loan-granting or hiring decision, one may use location of current residence as input. Even though location may not seem sensitive, it may be used to discriminate against people of a certain ethnicity, as on many occasions people from the same ethnic group congregate in certain spatially-focused communities [18]. Similar concerns exist for income or education level, which often exhibit strong correlations with the location where an individual works, lives or travels [9]. Note that, this sort of discrimination may often occur inadvertently, as opaque ML algorithms automatically exploit certain correlations in location data (e.g., higher default rates in certain zipcodes), without realizing their fairness implications.

(2) Fairness is achieved through some data transformation designed to prevent, or limit, the amount of bias in processing. This causes loss of utility, whereby the result of processing can be sub-optimal compared to the result obtained on the original data. Achieving fairness requires some utility loss, and the emerging fairness-utility trade-off must be carefully considered when devising a fairness mechanism. In the case of location data, utility has specific formulations, which may impact results in a way that is unique to spatial query processing algorithms. Using generic fairness mechanisms devised for other types of data may lead to poor utility, as seen in [27]. Therefore, it is desirable to design customized mechanisms for fighting bias in location data, such that the utility of spatial information is not significantly decreased.

In this paper, we introduce specific mechanisms targeted at providing *spatial fairness* while preserving data utility. We focus on the case of *individual fairness* [12], which is more difficult to achieve, but provides a higher level of fairness guarantees compared to its group-level counterpart. We provide specific definitions of location bias, and carefully characterize how location data can be used to exercise discriminatory decisions.

We introduce a novel construction called *fair polynomials* (Section 3.1) that can be used as building block within mechanisms for spatial fairness¹. We perform a detailed exploration of fair polynomials in order to understand their properties and the trade-off achieved between enforcing fairness and preserving data utility.

We identify two broad categories of scenarios where location bias occurs, and we define spatial fairness mechanisms for each:

- *Distance-based fairness* is relevant in location-based advertising and ride-hailing, where the dominant query type is

This work is licensed under the Creative Commons BY-NC-ND 4.0 International License. Visit <https://creativecommons.org/licenses/by-nc-nd/4.0/> to view a copy of this license. For any use beyond those covered by this license, obtain permission by emailing info@vldb.org. Copyright is held by the owner/author(s). Publication rights licensed to the VLDB Endowment.

Proceedings of the VLDB Endowment, Vol. 16, No. 2 ISSN 2150-8097.
doi:10.14778/3565816.3565820

¹While our focus is on geospatial data, some of our results can be extended for other types of data with continuous domains.

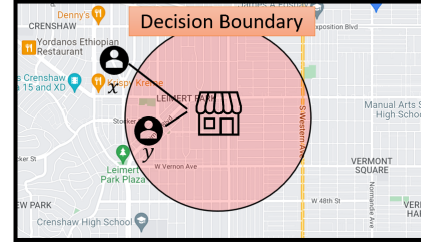
nearest-neighbors (NN). In this setting, location bias occurs when individuals are impacted by their distance to a reference point. In location-based marketing, an algorithm may advertise special deals to customers that are nearby a newly-opened health food store. The specific coordinates of the customers may be less relevant for utility, and instead the distance to a landmark is the important factor (i.e., the proximity to the store is a good indicator of likelihood of visiting it). While this may be efficient, it can have fairness repercussions. For instance, if the algorithm chooses the first 100 potential customers to reach based on distance, it is possible that only a rich neighborhood is covered. Customers from a poorer neighborhood that is adjacent to the rich one may never be selected, due to a slightly larger distance threshold. Figure 1a illustrates this case. In this situation, we would like to ensure that individuals from poorer backgrounds also get the chance to benefit from special deals on healthy food, even though they are slightly farther away (i.e., avoid the hard decision boundary phenomenon).

- *Zone-based fairness* is applicable in scenarios like gerrymandering, loan analysis or insurance pricing, where spatial range queries are the norm. In this case, we look at how to ensure spatial fairness with respect to coordinate values, instead of distances. This setting is broader, as it can provide fairness with respect to any reference point. Conversely, the amount of data utility sacrificed in the process may be higher. Figure 1b illustrates this point, where two individual homes are quoted significantly different insurance premiums due to their surrounding characteristics. Ideally, we would like the two residences to have similar premiums, given their physical proximity.

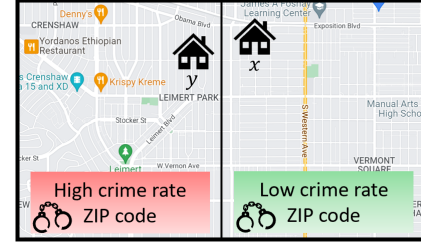
Our specific contributions are:

- We identify the problem of bias in spatial data processing, and formalize the notion of spatial data fairness;
- We introduce two definitions of spatial data fairness based on common interaction types, namely distance-based and zone-based fairness;
- We devise the novel concept of *fair polynomials*, which can be used as a building block to obtain mechanisms that achieve spatial fairness;
- We propose two mechanisms based on fair polynomials that enforce distance-based and zone-based fairness;
- We perform an extensive experimental evaluation on real datasets that shows the effectiveness of the proposed mechanisms, and investigates the fairness-utility trade-off.

The rest of the paper is organized as follows: Section 2 formalizes the notions of bias and fairness for location data. Sections 3 and 4 introduce our proposed distance-based and zone-based fairness mechanisms, respectively. We survey related work in Section 5. Section 6 reports the results of our experimental evaluation, followed by conclusions in Section 7.



(a) Hard decision-boundary on distance to landmark leads to unfairness towards poor neighborhoods.



(b) Despite close proximity, location x is assigned a much higher insurance premium compared to y .

Figure 1: Examples of two common location bias scenarios.

2 SYSTEM MODEL

2.1 Location Bias

The existence of bias, rooted in data or algorithms, is commonly used as a basis for reasoning on unfairness in decision-making. Several sources of bias have been identified in the literature, such as measurement bias [29] and behavioral bias [24], many of which are intertwined. In this work, we formalize a type of bias that occurs due to location data. Location bias is formally defined as follows:

Definition 1 (Location Bias). Distortion or algorithmic bias generated based on locations of entities in the geospatial domain or their distances to reference points is referred to as *location bias*.

Distortion refers to the bias intrinsic to the data. Whereas algorithmic bias [5] refers bias that is generated by processing algorithms. A category of bias closely related to algorithmic bias is called data processing bias [24], which occurs during data cleaning, enrichment, and aggregation. Our focus is on location bias sourced in processing algorithms. Location bias appears in a variety of applications where distances or locations may cause discrimination against individuals or groups. Consider the example in Fig. 2 in which a store wants to contact nearby customers with daily offers. A widely known algorithm such as *nearest neighbors* may be used to decide which customers should be contacted. If fairness measures are not considered, and if the store is located in a rich neighborhood, customers who live in less privileged areas of the city may never be contacted, and will be unable to enjoy the special offers.

Such an algorithmic advantage towards a particular group or towards certain individuals is an example of *distance-based* location bias with respect to a reference point. The distances can be represented as a one-dimensional input feature; however, the source of location bias can be multidimensional, more commonly stored in two or more dimensions, e.g., as latitudes and longitudes. Consider

classifying districts in a city as “high risk” or “low risk”, where more police presence and government resources are allocated to locations with higher crime rates. A strict boundary between a low-risk and high-risk district dictates that nearby individuals placed oppositely across the border are treated differently, despite their proximity. The unfairness manifests in the number of police patrols, the tolerance of police officers to crime, the cost of insurance, and the approval of home improvement loans. To address location bias in classification tasks, i.e., to prevent individuals with similar features from being treated significantly different, i.e., *unfairly*, we first introduce the notion of location fairness in Section 2.2, and then formulate the problem of achieving location fairness in Section 2.3.

2.2 Spatial Fairness Definition

There are two categories of fairness definitions, namely *group fairness* and *individual fairness* [13]. The former definition addresses the case where a group with certain features is treated statistically different compared to other groups. The latter approach focuses on treating individuals with similar features in a similar way. We adopt the use of individual fairness for spatial data, since (i) it provides higher fairness guarantees; (ii) it is more suitable for continuous domain features such as locations, in contrast to categorical features such as education, race, and gender; and (iii) location attributes tend to be dynamic, and hence more relevant on an individual basis, rather than for a group. To this end, we formally present the notion of *individual spatial fairness* in Definition 2:

Definition 2. (Individual Spatial Fairness). Let $\mathcal{L} = \{l_1, l_2, \dots, l_m\}$ denote the set of individual locations that need to be classified over the output set $\mathcal{A}$, where $\mathbf{l}_u \in \mathbb{R}^k$. A randomized mapping $M : \mathcal{L} \rightarrow \Delta(\mathcal{A})$ satisfies individual spatial fairness iff for every two locations $\mathbf{l}_u, \mathbf{l}_v \in \mathcal{L}$ the (D, d) -Lipschitz constraint holds,

$$D(M(\mathbf{l}_u), M(\mathbf{l}_v)) \leq d(\mathbf{l}_u, \mathbf{l}_v) \quad (1)$$

Intuitively, the definition states that the evaluation process M for two similar locations should yield similar outcomes. The definition relies on two key distance metrics, (1) similarity distance metric $d : V \times V \rightarrow \mathbb{R}$, measuring how similar individuals are, and (2) a distance metric $D(\cdot)$ measuring the distance between outcome distributions. The former metric will be defined thoroughly for locations in the upcoming sections, as it is tailored to the specific location interaction type. The latter metric, on the other hand, is commonly defined as *total variation norm* or so-called *statistical distance*. Given two probability distributions P and Q over outcome space $\mathcal{A}$, the statistical distance is calculated as

$$D(P, Q) = \frac{1}{2} \sum_{a \in \mathcal{A}} |P(a) - Q(a)|. \quad (2)$$

Our focus in this work is on binary decision-making tasks, hence, the output space is given by $\mathcal{A} = \{0, 1\}$. We assume a classifier modeled as a randomized mechanism $M : \mathcal{L} \rightarrow \Delta(\mathcal{A})$ mapping individuals over outcomes, where $\Delta(\mathcal{A})$ denotes all possible distributions. Thus, the classification of an individual $\mathbf{l}_u \in \mathcal{L}$ over outcome space $\mathcal{A}$ is done according to the distribution of $M(\mathbf{l}_u)$. To simplify notation, we assume function M to return the likelihood of the positive outcome, i.e., $M(\mathbf{l}_u) = M(\mathbf{l}_u|a=1)$.

Table 1: Summary of notations.

Symbol	Description
$\mathcal{L} = \{l_1, \dots, l_m\}$	Set of datapoints in $\mathbb{R}^k$
l_i	Distance from $\mathbf{l}_i$ to reference point
m	Number of datapoints
$\ \cdot\ _p$	p -norm distance
$\mathcal{A}$	Classification output domain
$d(\cdot)$	Distance between datapoints
$D(\cdot)$	Distance between distributions
$M(\mathbf{l}_i)$	Likelihood score of location $\mathbf{l}_i$

2.3 Problem Formulation

Consider m data points $\mathcal{L} = \{l_1, l_2, \dots, l_m\}$ located in a k -dimensional space ($\mathbb{R}^k$). Each location represents an individual. Associated attributes of data points are stored in a tabular format such that the i^{th} row of the table is dedicated to $\mathbf{l}_i$ (as shown in Fig. 2).

Attribute *Distance to Reference* (DtR) represents the distance from a reference point² $\mathbf{R}$. We use as DtR metric the *Minkowski* distance of order p (p -norm distance) defined for two data points $\mathbf{l}_u = (x_1, \dots, x_k)$ and $\mathbf{l}_v = (x'_1, \dots, x'_k)$ as

$$\|(\mathbf{l}_u, \mathbf{l}_v)\|_p = \sqrt[p]{\sum_{i=1}^k (x_i - x'_i)^p}. \quad (3)$$

The u -th entry of the DtR column is associated with datapoint $\mathbf{l}_u$ and entails a scalar $l_u = \|(\mathbf{l}_u, \mathbf{R})\|_p / \gamma$, where $\gamma = \max_{i=1 \dots m} \|(\mathbf{l}_i, \mathbf{R})\|_p$. Constant γ ensures the range of distances is $[0, 1]$ ($0 \leq l_i \leq 1, \forall i=1 \dots m$). The DtR column is shown with DtR metric set to 2-norm. It is important to note that DtR is based on data representation, and it is not to be confused with the two distance metrics $d(\cdot)$ and $D(\cdot)$, which are crucial elements of individual fairness.

As part of our system model, we assume a *classifier* performing decision-making on top of the data (this model aligns well with current location-based applications, where some machine learning is involved in data processing). Given an input individual u , the classifier $M(\mathbf{l}_u)$ returns a likelihood score for that individual based on her location (e.g., likelihood of receiving a location advertisement). Scores are real values in the range of 0 to 1.

The classifier output scores are shown in the last column of Fig. 2b. For example, user A is located at the coordinate $\mathbf{l}_A$ with the calculated distance from the reference point of $\sqrt{1^2 + 1^2} = \sqrt{2}$ normalized to $\sqrt{2}/\sqrt{10}$, and the generated score of $M(\mathbf{l}_A) = 0.8$. Other features and attributes used in the model could be race, gender, education, etc. The two problems we seek to address to achieve individual fairness are defined as follows:

PROBLEM 1. (Distance-based Spatial Fairness) *For a given location dataset $\mathcal{L}$ with the corresponding DtRs $\{l_1, \dots, l_m\}$, and a function $M : \mathcal{L} \rightarrow [0, 1]$, devise a mechanism to enforce individual distance-based fairness (D, d) -Lipschitz constraints with respect to DtRs.*

$$D(M(\mathbf{l}_i), M(\mathbf{l}_j)) \leq d(l_i, l_j) \quad \forall i, j \in [1, \dots, m] \quad (4)$$

PROBLEM 2. (Zone-based Spatial Fairness) *For a given location dataset $\mathcal{L} = \{l_1, l_2, \dots, l_m\}$, and a function $M : \mathcal{L} \rightarrow [0, 1]$, devise a*

²The proposed approach can be extended to multiple reference points by using a composite cost metric, or an existing locality-sensitive hashing (LSH) approach that produces a single scalar distance value over multiple reference points.

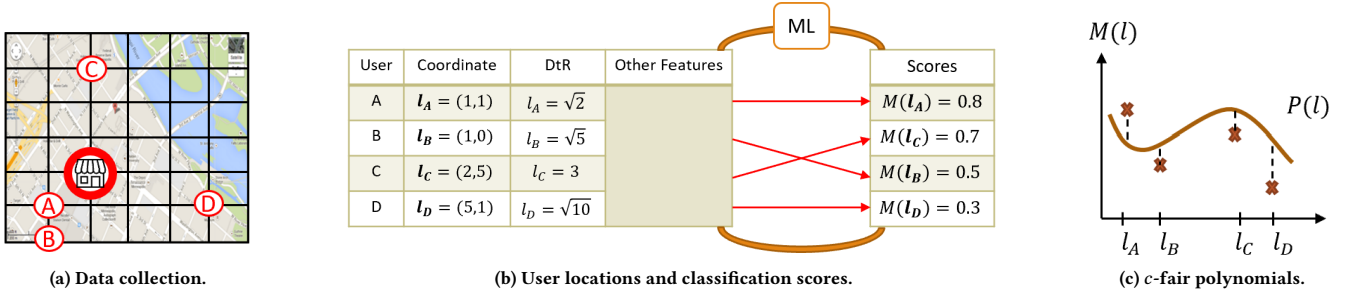


Figure 2: An example of distance-based fairness problem.

mechanism to enforce individual fairness (D, d)-Lipschitz constraints with respect to location coordinates.

$$D(M(l_i), M(l_j)) \leq d(l_i, l_j) \quad \forall i, j \in [1, \dots, m] \quad (5)$$

Fairness mechanisms must inherently alter the output likelihood scores in order to achieve the fairness requirement. Hence, there is a cost for such an operation in terms of utility loss. Since we expect the output of a fairness mechanism to be used for a learning task, we choose as utility metric *fitting error*, a widely accepted ML metric for output scores, formally presented in Definition 3.

Definition 3. (Utility). Let $\mathcal{B} : M \rightarrow M'$ be a mechanism that maps every likelihood score M to a likelihood score in M' given that $M, M' : \mathcal{L} \rightarrow \Delta(\mathcal{A})$. The fitting error (utility) of $\mathcal{B}$ is:

$$\sqrt{\frac{1}{m} \sum_{i=1}^m (M(l_i) - M'(l_i))^2} \quad (6)$$

As an example, suppose that the output score is mapped to a constant 0.5, i.e., $M'(l_i) = 0.5, \forall i$. The utility loss can then be calculated as $\sqrt{\frac{1}{4} (0.3^2 + 0.2^2 + 0.2^2)} \approx 0.206$.

3 DISTANCE-BASED SPATIAL FAIRNESS

We introduce a spatial fairness mechanism for Problem 1. Each data point is augmented with a DtR column representing a user's distance from the reference point. Therefore, the most natural similarity distance metric is 1-norm.

$$d(l_i, l_j) = |l_i - l_j| \quad (7)$$

LEMMA 3.1. Given the classifier output space of $\mathcal{A} = \{0, 1\}$, the statistical distance for every two individuals can be calculated as

$$D(M(l_i), M(l_j)) = |M(l_i) - M(l_j)| \quad (8)$$

PROOF.

$$D(M(l_i), M(l_j)) = \frac{1}{2} \sum_{a \in \{0,1\}} |P(a) - Q(a)| \quad (9)$$

$$= \frac{1}{2} (|M(l_i) - M(l_j)| + |1 - M(l_i) - (1 - M(l_j))|) \quad (10)$$

$$= |M(l_i) - M(l_j)| \quad (11)$$

□

3.1 Fair-Polynomials

Despite the strong fairness guarantees provided by individual fairness, applying a large number of hard constraints has limited its practicality. A common mechanism for individual fairness is to define an application-specific optimization problem usually referred to as *vendor's utility function* and solve it while imposing individual fairness hard constraints. Unfortunately, two major issues arise with such an approach when applied to location data: (i) In existing approaches, e.g., [27], a constraint optimization problem solver is used to alter input locations such that fairness requirements are met. The number of constraints grows quadratically with the number of data points, which makes their enforcement computationally prohibitive; (ii) The definition of the utility function in most scenarios is not straightforward, confining applicable use cases.

We devise the concept of *fair polynomials*, the intuition behind which is depicted in Figure 2c. A polynomial is efficiently fitted to the output scores of the classifier with a reasonably low *fitting error*. Fair polynomials no longer require enforcement of a large number of hard constraints: given a new data point, its corresponding fair value can be generated by evaluating the polynomial at that point.

Definition 4 (c -fair Polynomials). A single variable degree n polynomial $P(x) : \mathbb{R} \rightarrow \mathbb{R}$ is said to be c -fair if and only if for every two points x and y in its domain

$$|P(x) - P(y)| \leq c|x - y| \quad (12)$$

Given that a fair polynomial providing good estimates of likelihood scores exists, by one-to-one mapping (fitting) of the likelihood scores to polynomials, individual distance-based fairness can be achieved for every two data points. The constant $c \in [1, +\infty]$ in c -fair polynomials aims to exploit the trade-off between utility and fairness. When $c = 1$, the optimal individual location fairness is achieved but is usually associated with a higher loss in utility. When the value of c grows larger, the fairness constraint is relaxed, leading to higher utility but lower fairness.

As the distance-based fairness problem only involves scalars, our focus is on single variable degree n polynomials. We will extend to multi-variable polynomials to accommodate for multi-dimensional data points and address Problem 2 in Section 4.1. In the following, we answer three central questions, (i) what is the sufficient condition for a polynomial to be fair, (ii) how to derive the coefficients of the polynomial by imposing individual fairness constraints, and (iii) how to determine the degree of a fair polynomial.

3.2 Sufficient Condition for Fair Polynomials

There are several families of polynomials that preserve individual fairness over the defined distances for DtR. For example, one such family of polynomials is $P(x) = cx^n/n$ which is proven in Lemma 3.2 to be a c -fair polynomial.

LEMMA 3.2. *The polynomial $P(x) = cx^n/n$, is a c -fair polynomial for every two points $x, y \in [-1, 1]$.*

PROOF. The proof can be derived by expanding the equation and applying triangle inequality, considering that $|x^i y^j| \leq 1, \forall i, j$.

$$c \left| \frac{x^n - y^n}{n} \right| = c \left| \frac{(x - y)(x^{n-1} + x^{n-2}y + \dots + y^{n-1})}{n} \right| \leq \quad (13)$$

$$\frac{(c|x - y|)(|x^{n-1}| + |x^{n-2}y| + \dots + |y^{n-1}|)}{n} \leq c|x - y| \quad (14)$$

A fair polynomial must be flexible enough to reduce the error once the likelihood scores are fitted to the polynomial, and not every fair family of polynomials is a viable option. Consider the generic degree n polynomial written as

$$P(x) = a_0 + a_1x + \dots + a_nx^n, \quad (15)$$

where a_i are real numbers. In Theorem 1, we derive a sufficient condition for polynomials of order n to preserve individual fairness.

THEOREM 1. *A sufficient condition for a single variable degree n polynomial $P(x) = \sum_{i=0}^n a_i x^i$ to be c -fair given that $a_i \in \mathbb{R}$ and $|x| \leq 1$ is to have,*

$$\sum_{i=1}^n i|a_i| \leq c \quad (16)$$

PROOF. Following the definition of individual location fairness:

$$|P(x) - P(y)| = \left| \sum_{i=1}^n a_i (x^i - y^i) \right| \leq \sum_{i=1}^n |a_i| (x^i - y^i) \quad (17)$$

$$\leq \sum_{i=1}^n (|a_i| (x - y)) (|x^{i-1}| + |x^{i-2}y| + \dots + |y^{i-1}|) \quad (18)$$

$$\leq \sum_{i=1}^n |a_i| \times i(x - y) = |x - y| \sum_{i=1}^n i|a_i| \quad (19)$$

The above inequality is true based on Jensen's inequality (and also, extended triangle inequality) as well as applying the result from Lemma 3.2. Given that the inequality in Eq. (16) is satisfied, the polynomial is proven to be c -fair based on the definition. $\square$

The theorem indicates that if likelihood scores generated by the model are fitted to a polynomial for which the coefficients are selected such that $\sum_{i=1}^n i|a_i| \leq c$, then c -fairness is guaranteed for data entries. The sufficient condition in Theorem 1 can be used directly to learn c -fair polynomials, but the non-linearity existing in the constraint can result in higher computation complexity, as coefficients are unbounded. Theorem 2 addresses this problem by deriving linear constraints over coefficients.

THEOREM 2. *A sufficient condition for a 1-variable n -th degree polynomial $P(x) = \sum_{i=1}^n a_i x^i$ to be c -fair is to have:*

$$|a_i| \leq \frac{6 \times i \times c}{n(n+1)(2n+1)} \quad \forall i \in 1 \dots n \quad (a_i \in \mathbb{R}) \quad (20)$$

PROOF. The bound on each a_i value must allow for the maximum degree of freedom while fitting the likelihood scores. Therefore, the condition can be written as an optimization problem.

$$\begin{aligned} &\text{Minimize} && - \sum_{i=1}^n a_i \\ &\text{Subject to} && \sum_{i=1}^n i|a_i| \leq c \end{aligned} \quad (21)$$

Writing the Lagrangian and applying the stationary condition of Karush-Kuhn-Tucker (KKT) [15],

$$L(a_1, \dots, a_n, \lambda) = - \sum_{i=1}^n a_i - \lambda \left(\sum_{i=1}^n i|a_i| - c \right) \quad (22)$$

$$\Rightarrow \frac{\partial L}{\partial a_i} = -1 - \lambda \frac{i}{a_i} = 0 \rightarrow a_i = -\lambda i \quad (23)$$

λ can be derived from complementary slackness to be

$$\lambda \sum_{i=1}^n i^2 = c \rightarrow \lambda = \frac{6c}{n(n+1)(2n+1)} \quad (24)$$

Therefore, bounds on the coefficients are given as

$$|a_i| \leq \frac{6ic}{n(n+1)(2n+1)} \quad \forall i \in 1 \dots n \quad (25)$$

$\square$

3.3 Derivation of Fair Polynomials

We employ a simple ML model to compute polynomial coefficients, where each location distance represents a training sample used to fit the likelihood scores to a polynomial. The training set can be assembled by choosing at random a number of locations from the same data domain as the application (e.g., residential locations within a city). In cases where the target user population is already known (e.g., the coordinates of customers for a store that is using location-based advertising), this user set can be directly used for training. In the following, to simplify notation, we assume the latter. For a given training input l_i , the polynomial output is derived as

$$P(l_i) = a_0 + a_1 l_i + \dots + a_n l_i^n, \quad (26)$$

We denote the matrix of all training examples as

$$L = \begin{bmatrix} 1 & l_1 & l_1^2 & \dots & l_1^n \\ 1 & l_2 & l_2^2 & \dots & l_2^n \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & l_n & l_n^2 & \dots & l_n^n \end{bmatrix}.$$

Recall that m is the number of training examples, and the variables that we learn are the a_i s, which define the fair polynomial fitted to the data. The vector of coefficients can be written as

$$\mathbf{a}^T = [a_0, a_1, a_2, \dots, a_n] \quad (27)$$

and the likelihood scores are vectorized as

$$\mathbf{b}^T = [M(l_1), M(l_2), \dots, M(l_m)] \quad (28)$$

The convex optimization problem to learn $\mathbf{a}$ is formulated as:

$$\begin{aligned} &\text{Minimize} && \|\mathbf{L}\mathbf{a} - \mathbf{b}\|_2 \\ &\text{Subject to} && |a_i| \leq \frac{6 \times i \times c}{n(n+1)(2n+1)} \end{aligned} \quad (29)$$

This is equivalent to the least square problem with linear constraints and can be solved efficiently with algorithms such as Trust Region Reflective [11] and Bounded-variable least-squares [28], with complexity linear to the order of the polynomial n . Existing work [27] requires enforcing $O(m^2)$ hard constraints ($m \gg n$).

The selection of the polynomial degree n can be conducted based on a trial and error methodology. The optimal degree is the one that results in the minimum variance of error between likelihood scores and their corresponding values on the polynomial. Formally, let e_i denote the error between $M(\mathbf{l}_i)$ and $P(\mathbf{l}_i)$, i.e.,

$$e_i = |M(\mathbf{l}_i) - P(\mathbf{l}_i)| \quad (30)$$

Then, the value of $n \geq 1$ is selected such that

$$n = \operatorname{argmin}(\sum_{i=1}^m e_i^2) / (m - n - 1) \quad (31)$$

4 ZONE-BASED SPATIAL FAIRNESS

Revisiting the example in Figure 2, suppose that two users A and B both apply for a home improvement loan, and despite living in close proximity, one is categorized in an underdeveloped area and the other in a developed region due to geographic segmentation. A bank applies a classifier to decide whether an applicant should be granted a loan. The applicant whose home is in the underdeveloped category might be disadvantaged, as the location category can significantly impact the output of the classifier. Individual fairness argues that if two users are located close to each other, their output likelihood scores should not differ significantly.

In the distance-based fairness case, a single variable c -fair polynomial can fit output scores due to scalar distances. For multi-dimensional data points, Definition 4 is no longer directly applicable. To address this problem, we extend the definition to multivariate polynomials to achieve individual fairness for higher dimensional data points. The number of variables involved in fair polynomials is equal to the dimensionality of data points (k).

Three key variables are involved in finding an efficient family of fair polynomials that can fit the output likelihood scores with low utility loss: (i) dimensionality of data k ; (ii) the distance metric $d(\cdot)$ and (iii) fair polynomial degree n . The individual fairness problem can be characterized with respect to these criteria as follows:

- One-dimensional data representation (scalars), 1-norm distance, flexible order polynomial. This corresponds to the distance-based fairness case.
- 2-Dimensional data representation, 2-norm distances; order 1 polynomial. This is the most common scenario for locations where the attribute columns include 2D coordinates, and the fairness must be achieved with respect to Euclidean distance between individuals.
- k -Dimensional data representation, 2-norm distances; order 1 polynomial.
- k -Dimensional data representation, p -norm distances; order 1 polynomial.
- k -Dimensional data representation, p -norm distances; flexible order polynomial.

We formulate and derive the sufficiency condition to guarantee individual fairness for each mentioned scenario. The optimization problem in Eq. (29) is formulated with the derived constraints. We

omit the vectorization process for conciseness. For several of the proofs used in this section, we make use of Generalized Titu's Lemma provided in Lemma 4.1.

LEMMA 4.1 (GENERALIZED TITU'S LEMMA). *Let m be an integer greater than or equal to 2, a_i^m a non-negative real number, and x_i a positive real number. Then,*

$$n^{m-2} \sum_{i=1}^n \frac{a_i^m}{x_i} \geq \frac{(\sum_{i=1}^n a_i)^m}{\sum_{i=1}^n x_i} \quad (32)$$

PROOF. Proof is in Appendix B of our extended version³. $\square$

4.1 2-norm, 2 dimensional data, Order 1 polynomial

For higher dimensional data, the most common scenario happens when data points are in 2D, and the order of the polynomial is one. In practice, data points represent coordinates of locations on the map. Consider two locations $\mathbf{l}_1 = (x_1, x_2)$ and $\mathbf{l}_2 = (x'_1, x'_2)$ in $\mathbb{R}^2$, where x_1 and x'_1 are the x -axis coordinates, while x_2 and x'_2 denote y -axis coordinates. To achieve individual fairness with respect to locations, the hard Lipschitz constraints dictate that:

$$D(M(\mathbf{l}_i), M(\mathbf{l}_j)) \leq d(\mathbf{l}_i, \mathbf{l}_j) \quad \forall i, j \in 1 \dots m \quad (33)$$

The distance between distribution scores, i.e., $D(\cdot)$, is calculated as before based on Equation (3.1) and the distance between locations is the 2-norm of data points (Euclidean distance), calculated as:

$$d(\mathbf{l}_i, \mathbf{l}_j) = \sqrt{(x_1 - x'_1)^2 + (x_2 - x'_2)^2} \quad (34)$$

We start by showing how a fair-polynomial can be derived for the Euclidean similarity distance. Then, we relax the assumptions and generalize the approach for arbitrary distance norms as well as n -dimensional data points. Recall that as location data are stored in 2D, the fair polynomial consists of two variables. The generalized definition of fair polynomials for order n polynomials and k dimensional data is provided in Definition 5.

Definition 5 (Generalized c -Fair Polynomial). The polynomial $P(x_1, x_2, \dots, x_k) : \mathbb{R}^k \rightarrow \mathbb{R}$ with real coefficients is c -fair iff for every two points $\mathbf{x} = (x_1, x_2, \dots, x_m)$ and $\mathbf{x}' = (x'_1, x'_2, \dots, x'_m)$ in its domain

$$|P(x_1, x_2, \dots, x_m) - P(x'_1, x'_2, \dots, x'_m)| \leq c \times d(\mathbf{x}, \mathbf{x}') \quad (35)$$

In the case of 2-dimensional locations and Euclidean distance, fair polynomials imply that for every two locations $\mathbf{l}_1 = (x_1, x_2)$ and $\mathbf{l}_2 = (x'_1, x'_2)$, we must have,

$$|P(x_1, x_2) - P(x'_1, x'_2)| \leq c \times \sqrt{(x_1 - x'_1)^2 + (x_2 - x'_2)^2} \quad (36)$$

Where the polynomial is denoted by

$$P(x_1, x_2) = a_0 + a_1 x_1 + a_2 x_2 \quad (37)$$

The goal is to learn the coefficients a_i such that the polynomial $P(\cdot)$ can model the output scores $M(\cdot)$ and preserve fairness with respect to Euclidean distance. Theorem 3 provides the sufficiency condition for a two-variables order one polynomial to be fair.

³<https://github.com/SinaShaham/c-Fair-Polynomials/tree/main>

THEOREM 3. A sufficient condition for a 2-variable first degree polynomial $P(x_1, x_2) = a_0 + a_1x_1 + a_2x_2$ to be c -fair defined over 2-norm similarity distance is to have:

$$|a_1|, |a_2| \leq c/\sqrt{2} \quad (a_1, a_2 \in \mathbb{R}) \quad (38)$$

PROOF. On the one hand, based on Lemma 4.1, a lower bound for Euclidean distances can be written as

$$d(\mathbf{l}_i, \mathbf{l}_j) = \sqrt{(x_1 - x'_1)^2 + (x_2 - x'_2)^2} \geq \quad (39)$$

$$\sqrt{(|x_1 - x'_1| + |x_2 - x'_2|)^2/2} \quad (40)$$

On the other hand, for the polynomial one can write

$$|P(x_1, x_2) - P(x'_1, x'_2)| = |a_1(x_1 - x'_1) + a_2(x_2 - x'_2)| \quad (41)$$

$$\leq |a_1|(|x_1 - x'_1|) + |a_2|(|x_2 - x'_2|) \quad (42)$$

By combining the two equations the sufficiency condition in Equation (38) can be derived from the following inequality

$$|a_1|(|x_1 - x'_1|) + |a_2|(|x_2 - x'_2|) \leq c \times (|x_1 - x'_1| + |x_2 - x'_2|)/\sqrt{2} \quad (43)$$

□

The above theorem indicates that if the coefficients of polynomials fitted to data are chosen such that $|a_1|, |a_2| \leq c/\sqrt{2}$, fairness is guaranteed for every two locations in the domain. The sufficiency condition for first degree polynomials is generalized for k -dimensional data points in space in Theorem 4 (the number of variables in the polynomial is equal to the number of dimensions).

THEOREM 4. A sufficient condition for a k -variable first degree polynomial $P(x_1, \dots, x_k) = a_0 + \sum_{i=1}^k a_i x_i$ defined over 2-norm similarity distance to be c -fair is:

$$|a_i| \leq c/\sqrt[2]{k}, \quad \forall i = 1 \dots k, \quad (a_i \in \mathbb{R}) \quad (44)$$

PROOF. Please see proof in Appendix B of our extended version.

□

4.2 p -norm, k dimensional, Order 1 polynomial

We relax the similarity metric for arbitrary p -norm distance, calculated for data points $\mathbf{l}_i = (x_1, \dots, x_k)$ and $\mathbf{l}_j = (x'_1, \dots, x'_k)$ as

$$d(\mathbf{l}_i, \mathbf{l}_j) = \sqrt[p]{\sum_{q=1}^k (x_q - x'_q)^p} \quad (45)$$

THEOREM 5. A sufficient condition for a k -variable first degree polynomial $P(x_1, \dots, x_k) = a_0 + \sum_{i=1}^k a_i x_i$ defined over p -norm similarity distance to be c -fair is:

$$|a_i| \leq c/\sqrt[p]{k^{p-1}}, \quad \forall i = 1 \dots k \quad (46)$$

PROOF. Based on generalized Titu's Lemma, we have on the one hand a lower bound for Euclidean distances:

$$d(\mathbf{l}_i, \mathbf{l}_j) = \sqrt[p]{\sum_{q=1}^k (x_q - x'_q)^p} \geq \sqrt[p]{\left(\sum_{q=1}^k |x_q - x'_q|\right)^p / k^{p-1}} = \left(\sum_{q=1}^k |x_q - x'_q|\right) / \sqrt[p]{k^{p-1}} \quad (47)$$

On the other hand, for the polynomial one can write

$$|P(x_1, \dots, x_k) - P(x'_1, \dots, x'_k)| = \left| \sum_{q=1}^k a_q (x_q - x'_q) \right| \quad (48)$$

$$\leq \sum_{q=1}^k |a_q| |x_q - x'_q| \quad (49)$$

Combining the two, we obtain:

$$\sum_{q=1}^k |a_q| |x_q - x'_q| \leq \sum_{q=1}^k |x_q - x'_q| / \sqrt[p]{k^{p-1}} \quad (50)$$

The inequality is satisfied when $|a_q| \leq 1/\sqrt[p]{k^{p-1}}$. □

4.3 p -norm, k dimensional, Order n polynomial

So far, the sufficiency condition for c -fair polynomials was derived for arbitrary norms in k -dimensional space based on order 1 polynomials. Moreover, for distance-based fairness, c -fair polynomials were developed for 1-dimensional distance using arbitrary degree n polynomial. This subsection provides the theoretical background for the generalized scenario in which the location data are in k dimensions with the norm set to p , and degree n polynomials.

Although by increasing the degree of polynomials, a better fit to likelihood scores can be achieved, the existence of monomials in which multiple variables are involved leads to complexity in the derivation of sufficiency conditions. To address this, we assume that the monomials in the multivariable polynomial consist of only a single variable. Making such an assumption comes with the cost of utility loss; however, it greatly reduces the complexity of the generic case. We assume that the degree n polynomial is expressed as the summation of k univariate polynomials.

$$P(x_1, x_2, \dots, x_k) = \sum_{i=1}^k P_i(x_i), \quad (51)$$

where $P_i(x_i) = \sum_{j=1}^n a_{ij} x_i^j$ is a degree n univariate polynomial with its input being x_i , the i th variable in the original polynomial. The assumption helps to remove existence of monomials with multiple variables such as $x_i^3 x_j^2 x_k^3$ and to simplify derivation of location fairness sufficiency conditions provided in Theorem 6.

THEOREM 6. A sufficient condition for a k -variable n -th degree polynomial $P(x_1, x_2, \dots, x_k) = \sum_{i=1}^k P_i(x_i)$ to be c -fair defined over p -norm similarity distance is to have:

$$|a_{ij}| \leq \frac{6 \times j \times c \sqrt[p]{k^{p-1}}}{n(n+1)(2n+1)}, \quad \forall i = 1 \dots k \ \& \ \forall j = 1 \dots n \quad (52)$$

PROOF. We write the equation in its component form shown in Equation 51. An upper bound for $P_i(x_i)$ can be derived as,

$$|P(x_1, \dots, x_k) - P(x'_1, \dots, x'_k)| = \left| \sum_{i=1}^k P_i(x_i) - \sum_{i=1}^k P_i(x'_i) \right| \quad (53)$$

$$= \left| \sum_{i=1}^k (P_i(x_i) - P_i(x'_i)) \right| \leq \sum_{i=1}^k |P_i(x_i) - P_i(x'_i)| \quad (54)$$

An upper bound for the sub-terms for all $j = 1 \dots k$ can be derived as

$$|P_i(x_i) - P_i(x'_i)| = \left| \sum_{j=1}^n a_{ij}(x_i^j - x_i'^j) \right| \leq |x_i - x'_i| \sum_{j=1}^n j|a_{ij}| \quad (55)$$

The above inequality is derived based on Eq. (17) setting $|x_i| \leq 1$. We also use the lower bound derived in Eq. (47)

$$c \times d(\mathbf{l}_u, \mathbf{l}_v) \geq c \sum_{q=1}^k |x_q - x'_q| / \sqrt[k]{k^{p-1}} \quad (56)$$

Putting the derived upper bound and lower bound together, the following inequality is satisfied, and c -fairness is guaranteed.

$$|x_i - x'_i| \left(\sum_{j=1}^n j|a_{ij}| \right) \leq c|x_j - x'_j| / \sqrt[k]{k^{p-1}} \quad (57)$$

$$\sum_{j=1}^n j|a_{ij}| \leq c / \sqrt[k]{k^{p-1}} \quad (58)$$

By applying the method used in DtR, the bounds are linearized to,

$$|a_{ij}| \leq \frac{6 \times j \times c}{n(n+1)(2n+1) \sqrt[k]{k^{p-1}}} \quad \forall i \in 1 \dots n \quad (a_i \in \mathbb{R}) \quad (59)$$

□

It is worth noting that the generated scores by fair polynomials can result in values greater than one or less than zero. In such scenarios, the values are suppressed to one and zero, respectively. It is straightforward to prove that the suppression process does not violate the individual fairness constraints and leads to higher utility.

5 RELATED WORK

Fairness notions can be grouped into two broad categories of Group Fairness and Individual Fairness [12]. In group fairness, a protected attribute of the dataset, such as race or gender, which is considered to be critical in decision-making outcomes, partitions individuals into groups. The ML model used for a decision-making task on the dataset is considered to be fair if it achieves some statistical measure across groups. A few of the key statistical measures include statistical parity [21][12], equalized odds [17], treatment equality [7], and test fairness [10]. Individual fairness aims to give similar predictions to similar outcomes, focusing on fairness for individuals as opposed to groups. Group fairness notions are generally weaker than individual fairness notions [19]. Despite higher fairness guarantees provided by individual fairness and fragility of group fairness notions, group fairness notions are widely studied in the literature due to their easier enforcement [23]. Only a handful of approaches exist in the literature to achieve fairness in the geospatial domain.

The current state-of-the-art approach to enforce individual location fairness is to define a linear loss function once the likelihood scores are generated and solve optimization under individual fairness Lipschitz constraints. Let I be an instance of our problem consisting of a metric $d : \mathcal{L} \times \mathcal{L} \rightarrow \mathbb{R}$, and a loss function $L : \mathcal{L} \times A \rightarrow \mathbb{R}$, the optimization problem is defined as,

$$\text{opt}(I) = \min_{\{M(x)\}_{x \in L}} \mathbb{E}_{x \sim L} \mathbb{E}_{a \sim M(x)} L(x, a) \quad (60)$$

$$\text{Subject to: } \forall x, y \in L : D(M(x), M(y)) \leq d(x, y) \quad (61)$$

$$\forall x \in L : M(x) \in \Delta(A) \quad (62)$$

One can see that the number of constraints in this mechanism grows quadratically with the number of individuals, imposing a large computational complexity on the system. The authors in [27] formulate the loss function for location-based advertisements in social media. Locations visited on the map are shown as binary strings, and a classifier is used to predict whether a user should receive a targeted advertisement. Moreover, not directly related to locations, but for general purpose advertisement and auctions, individual fairness is applied in [13]. Another application over which the loss function has been defined is achieving individual fairness in ranking and recommendation systems [26]. In ranking systems, the amount of unfairness with respect to individuals is measured after ranking, and a loss function aims to reorder ranking such that the amount of individual unfairness is minimized [8].

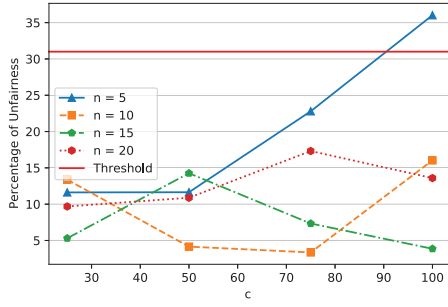
Several attempts have also been made to apply the individual fairness notion for clustering datapoints in Cartesian space. The notion in [20] defines clustering conducted for a point in space as fair if the average distance to the points in its own cluster is not greater than the average distance to the points in any other cluster. The authors in [22] focus on defining individual fairness for k -median and k -means algorithms. Clustering is defined to be individually fair if every point expects to have a cluster center within a particular radius. To the best of our knowledge, no work has directly defined individual fairness with respect to locations.

6 EXPERIMENTAL EVALUATION

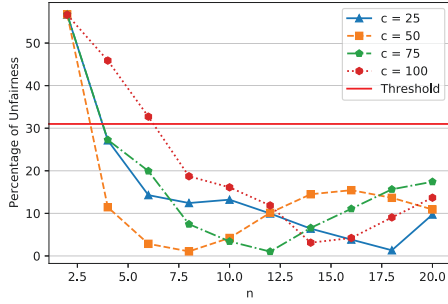
We evaluate our proposed spatial data fairness mechanisms in the two studied scenarios. For the distance-based case, we use a dataset of taxi fares from New York City; for the zone-based case, we consider budget allocation to police departments according to the Chicago crime occurrence dataset. We ran experiments on a 3.40GHz core-i7 Intel processor with 8GB RAM. The code is implemented in Python and uses the Trust Region Reflective least square implementation from Scipy [4] (the maximum number of iterations for convergence is set to 300, and the default tolerance threshold value of $1e-2$ is used to stop the optimization iterations).

6.1 Distance-based Spatial Fairness

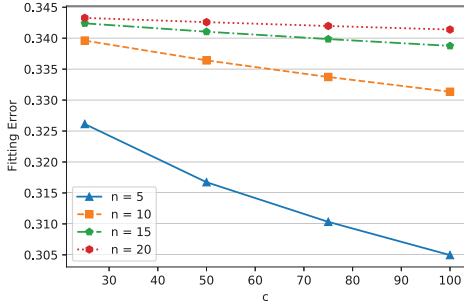
We sampled 120,000 records from the NYC taxi dataset [1] providing over 55 million trips and their associated fares. We deployed an Artificial Neural Network (ANN) to assess the likelihood of taxi fares being fair in the system. Specifically, we seek to capture whether there is bias in the setting of fares based on the specific neighborhoods where the trip starts/ends. Ideally, the trip distance should be the



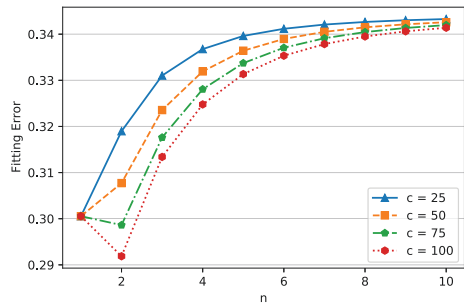
(a) Percentage of Unfairness versus c .



(b) Percentage of Unfairness versus n .



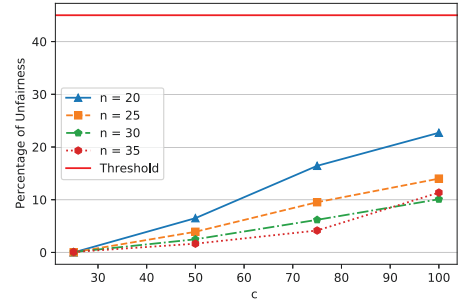
(c) Fitting Error versus c .



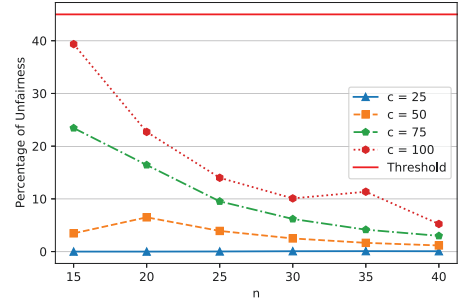
(d) Fitting Error versus n .

Figure 3: Distance-based Mechanism, New York taxi dataset.

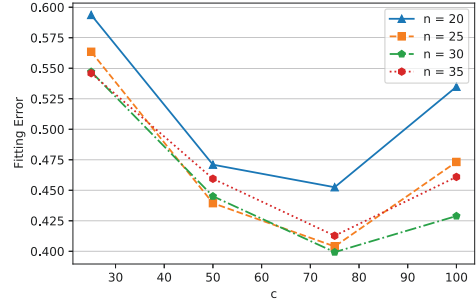
only factor determining price (we carefully pre-process the dataset such that trips are clustered according to the time of week/day, such that differences in fare due to demand status and traffic causes are eliminated). Our goal is to first understand the percentage of records for which the individual fairness constraints do not hold with respect to traveled distances. Then, we analyze the performance of the proposed c -fair mechanism.



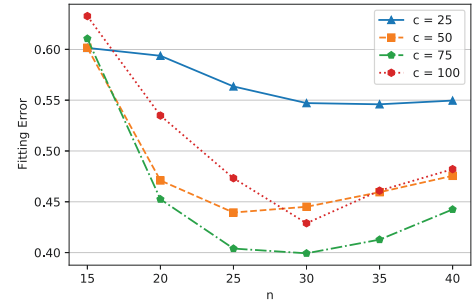
(a) Percentage of Unfairness versus c .



(b) Percentage of Unfairness versus n .



(c) Fitting Error versus c .



(d) Fitting Error versus n .

Figure 4: Zone-based Mechanism.

ML Model for Fairness Characterization. Our ANN model consists of two hidden layers with 200 and 100 neurons and an output layer with two neurons representing the binary classification task. The activation function used in the model is RELU, the dropout probability for each layer set to 0.4, and cross-entropy is used as the loss function. The accuracy of the model is 92%. The input features include pick-up date and time (categorical hour, AM or PM, weekday,

EDT date), pick-up longitude, pick-up latitude, drop-off longitude, drop-off latitude, passenger count, and distance traveled in kilometers. The ride fares have a mean of 10 dollars with a standard deviation of 7 dollars, and the average traveled distance is 3.31 km with a standard deviation of 3.2 km. For model training purposes, we split the data into training, validation and test datasets with 96,000, 12,000, and 12,000 records, respectively. To generate the ground truth for the training dataset, we have used price per kilometer traveled as the indicator of how fair the associated traveled fares are. For every hour of the day, the average price per kilometer is calculated as the hard threshold between fair and unfair travels. The trips above the threshold are classified as unfair, and the values less than the threshold are assumed to be fair. This results in a total of 21,928 trips being categorized as unfair.

Once the ANN model is trained, we predict the likelihood of each trip fare being fair on the test dataset. For every two records, the individual fairness constraint is evaluated to reveal whether fares are fair with respect to travel distances. In the absence of any fairness mechanisms being deployed, 32% of constraints are not satisfied, hence those trips are unfair. The 32% threshold is highlighted with a red horizontal line in the experimental plots, to highlight the fairness improvement of the considered approaches.

Next, we apply the proposed c -fair mechanism to achieve distance-based fairness. Our experiments evaluate the performance based on four key metrics: percentage of unfairness (constraints were not satisfied), the degree of c -fair polynomial (n), the parameter c , and the root mean square (RMS) of fitting error to likelihood scores.

Percentage of Unfairness. Fig. 3a shows the impact of increasing c on reducing unfairness when the degree of the polynomial is 5, 10, 15, and 20. As expected, lower values of c result in higher fairness in the system, with maximum fairness achieved when c is equal to one. For the maximum fairness scenario, the percentage of unfairness is zero, meaning that all individual fairness constraints are satisfied for every two records in the dataset. By increasing the value of c , fair polynomials would have more room for maneuver and fitting to likelihood scores, but it comes with the cost of higher unfairness. Such behaviour demonstrates the utility-fairness trade-off captured by the constant c . Increasing the polynomial degree can be seen to improve the percentage of unfairness until it reaches the point where it overfits the likelihood scores, and the performance deteriorates. Fig. 3b shows more clearly the impact of increasing the value of n on unfairness. Lower c values result in a lower percentage of unfairness for all polynomial degrees.

Fitting Error. Figs. 3c and 3d demonstrate the amount of utility loss in data due to fitting likelihood scores to a c -fair polynomial. Two key trends can be observed from the figures. First, increasing the value of c lowers the fitting error. This is expected, as higher c allows more flexibility for selecting coefficients and better fitting performance. Second, increasing the value of n for the same value of c raises the fitting error. To understand this behavior, one can intuitively look at the problem as allocating the same amount of budget among several buckets representing coefficients. Although increasing the degrees of freedom provides better fitting performance as higher degree monomials exist, it further restricts the budget for each coefficient. Thus, the lower degree monomials, which have a more significant impact on the performance, are allocated a lower amount of budget, negatively affecting the performance.

Computational Complexity. We measure the computational overhead of c -fair polynomials in terms of time complexity, number of iterations before optimization convergence, and final optimization cost (the final optimization cost represents the value of the Scipy cost function upon reaching the solution [4]). Fig. 5 shows the results. In each graph, the overhead is shown for four values of $c = 25, 50, 75, 100$ plotted for varying polynomial degrees. Overall, the time complexity is in the order of milliseconds and does not limit the practical deployment of c -fair polynomials. The second graph illustrates the number of iterations before reaching the optimal point. The optimization process stops either by reaching the maximum number of iterations (300) or when the relative change in optimization cost remains below the tolerance threshold ($1e-2$). As explained previously, the slight oscillation in the performance is due to selecting a random start point for the optimization.

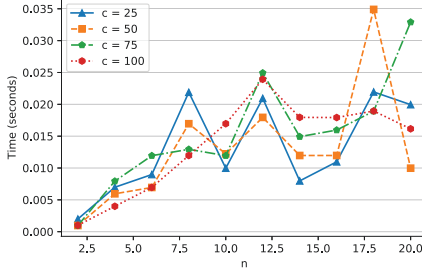
Note that, increasing the degree of polynomial n results in higher computational overhead. This is expected, as more degrees of freedom (coefficients) lead to more effort for finding the optimal point. Another consistent behavior across all three figures is that increasing c on average reduces the computation complexity cost and facilitates reaching the near-optimal point. The trend is more apparent in the final optimization cost figure, in which it can be clearly seen that a higher c value leads to a lower cost.

6.2 Zone-based Spatial Fairness

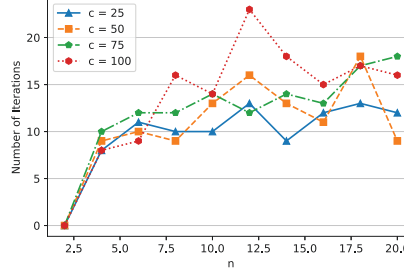
For this scenario, we consider the case of budget allocation to different areas of Chicago, USA, based on the measured crime rates. We use the dataset provided by the Chicago Police Department's CLEAR (Citizen Law Enforcement Analysis and Reporting) system [3], consisting of reported crime incidents in Chicago. A 1024×1024 grid is overlaid on top of the Chicago map, and the goal is to fairly allocate the budget such that neighborhoods that are close to each other are treated similarly. We have selected seven major crime categories of sexual assault, homicide, kidnapping, sex offense, motor vehicle theft, criminal damage, and narcotics among the reported crimes, and trained a logistic regression model to infer the likelihood of crime occurrence in each cell. The training dataset includes the crime data from January to November 2015, and the December data is chosen as the test dataset. The accuracy of the model is 94% and its output is a set of likelihoods indicating the probability of crime occurrence. The budget allocated to each cell is proportional to the likelihood score derived by the classifier.

Once the likelihood scores are generated, they are used with X and Y cell coordinates to achieve individual location fairness with the distance metric set to 2-norm. In absence of any fairness mechanism, we determine the percentage of individual location fairness constraints not being satisfied at 44.0%.

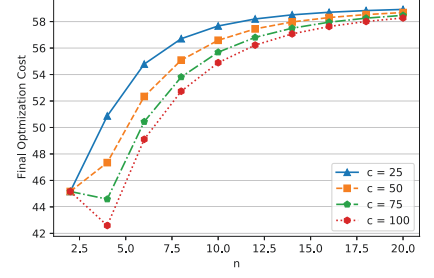
To understand if the expansion to higher polynomial degrees is essential, we started our experiments by focusing on degree one c -fair polynomials and applying the results in Theorem 3. As expected, the fitting error was rather high, and the utility was insufficient. We also noticed that for degree one polynomials, the optimal solution is achieved even when the value of c is equal to one. Therefore, increasing c does not help with improving the fitting error. Thus, it is crucial to use higher degree polynomials for this purpose.



(a) Time Complexity.

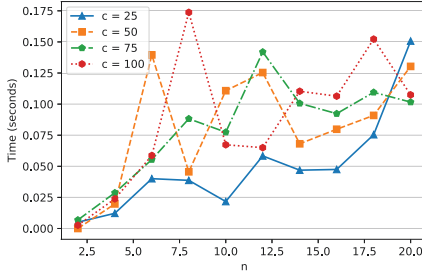


(b) Number of Iterations Before Convergence.

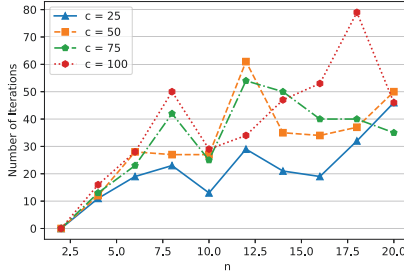


(c) Optimization Final Cost.

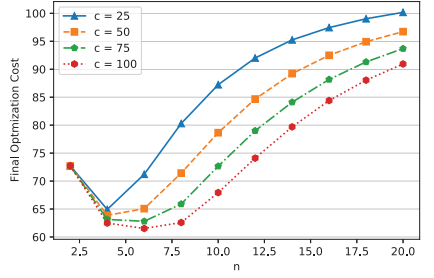
Figure 5: Computation Overhead Analysis on New York Taxi Dataset.



(a) Time Complexity.

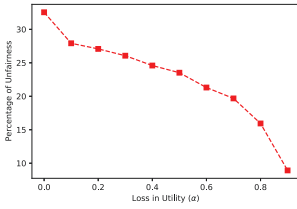


(b) Number of Iterations Before Convergence.

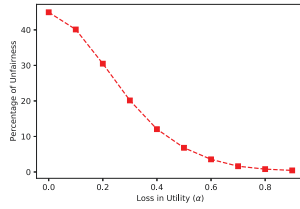


(c) Optimization Final Cost.

Figure 6: Computation Overhead Analysis on Chicago Crime Dataset.



(a) Distance-based Mechanism.



(b) Zone-based Mechanism.

Figure 7: Comparison Benchmark Performance.

Next, we apply the optimization formulation derived in Theorem 6, the most generalized formula allowing each dimension to contribute in fitting with a degree n polynomial. Fig. 4 demonstrates the performance of c -fair polynomials for achieving individual location fairness on the crime dataset. The patterns are generally consistent with the distance-based case considered in the New York taxi fares experiment. Figs. 4a and 4b show the impact of c and n on the percentage of unfairness and Figs. 4c and 4d show the performance with respect to utility. The red line is used as the reference point representing the percentage of unfairness in the original data, in the absence of fairness mechanisms.

Increasing the value of c results in a lower degree of fairness and higher fitting error once the degree of polynomial reaches an acceptable level. This result further substantiates the fairness-utility trade-off in the system, also observed in the distance-based case. In a 2-dimensional space, using a degree of 10 and the above polynomial to model each dimension can ensure that scores are under-fitted, preventing high fitting error values. In summary, the amount of fairness achieved with the fair polynomials, even for a

reasonably low degree for polynomials such as 15 and values of c greater than 10, can be seen to be over 70%.

Fig. 6 shows the computation complexity of the proposed mechanism. The first point to notice is that a relatively high amount of time is required to achieve zone fairness compared to the distance-based setting. However, computational complexity is still low in absolute value, and not an obstacle for practical deployment, with sub-second execution time.

6.3 Comparison with benchmarks

We derive a threshold-based benchmark from the fairness mechanism in [12] and the binary evaluation approach in [14]. Given constant threshold t , let polynomial $P(l_i) = t$, and let parameter α define by how much scores can be altered (e.g., if $\alpha = 0.1$, each score can be altered by at most ± 0.1). For DtR, the benchmark cycles through each score M_i and pushes it towards polynomial $P(l_i)$ with an allowance of α :

$$M(l_i) \leftarrow M(l_i) + \text{sign}(P(l_i) - \alpha) \times \min(\alpha, |P(l_i) - \alpha|) \quad (63)$$

The *sign* operation ensures the direction of change is favorable to the benchmark, and *min* ensures that the change in score does not overshoot $P(l_i)$. When α is zero, no utility loss exists, and the unfairness percentage is the same as the original. As α grows, the flexibility margin results in more constraints being satisfied until absolute fairness is achieved. For the zone-based fairness case, we use $P(x, y) = (x + y)/\sqrt{2}$, easily extensible to higher dimensions.

Figure 7 shows the utility-fairness trade-off obtained by the baseline for both distance- and zone-based cases. Unfairness is monotonically decreasing as one allows a higher α allocation to the score. We notice a plateau of desirable behavior concentrated

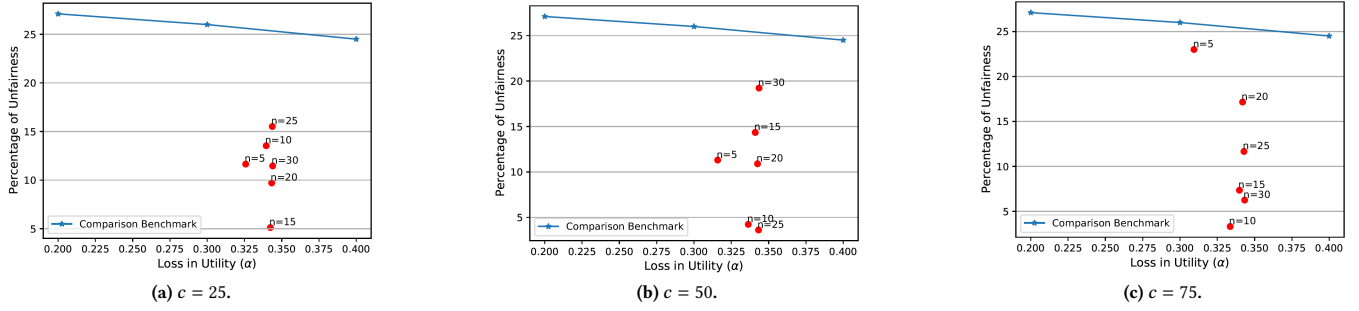


Figure 8: Distance-based Mechanism Benchmark Comparison, New York taxi dataset.

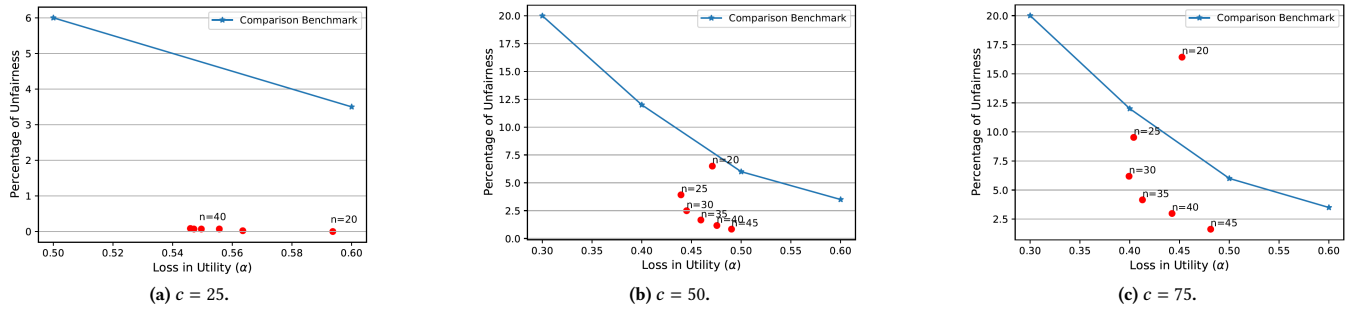


Figure 9: Zone-based Mechanism Benchmark Comparison, Chicago crime dataset.

between $\alpha = [0.2, 0.6]$, where the utility-fairness trade-off is good (for α values below 0.2 unfairness is too high, whereas for values higher than 0.6 the fairness gain fades in the multi-variate case).

Next, we focus our attention to this desirable α range, and we plot the relative performance of the benchmark versus our proposed fair polynomials approach for various values of c and n . The plots in Figs. 8 and 9 show that our approach provides a superior trade-off compared to benchmarks. The benchmark outperforms only for zone-based fairness when $n = 20$. In all other cases, fair polynomials provide either a vastly improved utility, or better fairness.

7 CONCLUSION

We studied in depth the problem of individual fairness for location data, and we identified sources of location bias that can occur in practical settings. We formulated two distinct settings for spatial fairness: distance-based and zone-based fairness, and we devised specific techniques to achieve spatial fairness while preserving utility, with the help of a novel construction called fair polynomials. While our focus is on spatial fairness, fair polynomials have the potential to provide useful building blocks for fairness in other application domains. In future work, we plan to study more complex types of spatial queries. At the same time, we plan to study the effect of fairness mechanisms in conjunction with other constraints, such as privacy, e.g., devise mechanisms that can achieve both spatial fairness and location privacy.

ACKNOWLEDGMENTS

This research has been funded in part by NSF grants IIS-1910950, IIS-1909806, CNS-2125530, NIH grant R01LM014026, the USC Integrated Media Systems Center (IMSC), and unrestricted cash gift from Microsoft Research and Google. Any opinions, findings, and

conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of any of the sponsors such as the NSF.

REFERENCES

- [1] [n.d.]. New York City taxi fare prediction. <https://www.kaggle.com/c/new-york-city-taxi-fare-prediction/data>
- [2] 2013. Chicago Metropolitan Agency for Planning. Travel patterns in economically disconnected area clusters. <https://www.cmap.illinois.gov/2050/maps/transit-edu>
- [3] 2015. Chicago Crime Dataset. <https://data.cityofchicago.org/Public-Safety/Crimes-2015/vwvp-7yr9>
- [4] 2021. SciPy Documentation Version: 1.7.1 (lsq linear). <https://docs.scipy.org/doc/scipy/>
- [5] Ricardo Baeza-Yates. 2018. Bias on the web. *Commun. ACM* 61, 6 (2018), 54–61.
- [6] Solon Barocas, Moritz Hardt, and Arvind Narayanan. 2017. Fairness in machine learning. *NIPS tutorial* 1 (2017), 2.
- [7] Richard Berk, Hoda Heidari, Shahin Jabbari, Michael Kearns, and Aaron Roth. 2021. Fairness in criminal justice risk assessments: The state of the art. *Sociological Methods & Research* 50, 1 (2021), 3–44.
- [8] Asia J Biega, Krishna P Gummadi, and Gerhard Weikum. 2018. Equity of attention: Amortizing individual fairness in rankings. In *The 41st international acm sigir conference on research & development in information retrieval*. 405–414.
- [9] Enrico Cantoni. 2020. A precinct too far: Turnout and voting costs. *American Economic Journal: Applied Economics* 12, 1 (2020), 61–85.
- [10] Alexandra Chouldechova. 2017. Fair prediction with disparate impact: A study of bias in recidivism prediction instruments. *Big data* 5, 2 (2017), 153–163.
- [11] Andrew R Conn, Nicholas IM Gould, and Philippe L Toint. 2000. *Trust region methods*. SIAM.
- [12] Cynthia Dwork, Moritz Hardt, Toniann Pitassi, Omer Reingold, and Richard Zemel. 2012. Fairness through awareness. In *Proceedings of the 3rd innovations in theoretical computer science conference*. 214–226.
- [13] Cynthia Dwork and Christina Ilvento. 2018. Individual fairness under composition. *Proceedings of Fairness, Accountability, Transparency in Machine Learning* (2018).
- [14] Alessandro Fabris, Stefano Messina, Gianmaria Silvello, and Gian Antonio Susto. 2022. Algorithmic Fairness Datasets: the Story so Far. *arXiv preprint arXiv:2202.01711* (2022).

- [15] Benyamin Ghogh, Ali Ghodsi, Fakhri Karray, and Mark Crowley. 2021. KKT Conditions, First-Order and Second-Order Optimization, and Distributed Optimization: Tutorial and Survey. *arXiv preprint arXiv:2110.01858* (2021).
- [16] Sara Hajian, Francesco Bonchi, and Carlos Castillo. 2016. Algorithmic bias: From discrimination discovery to fairness-aware data mining. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*. 2125–2126.
- [17] Moritz Hardt, Eric Price, and Nati Srebro. 2016. Equality of opportunity in supervised learning. *Advances in neural information processing systems* 29 (2016), 3315–3323.
- [18] Kamyoun Kim. 2020. A Spatial Optimization Approach for Simultaneously Districting Precincts and Locating Polling Places. *ISPRS International Journal of Geo-Information* 9, 5 (2020), 301.
- [19] Michael P Kim, Aleksandra Korolova, Guy N Rothblum, and Gal Yona. 2019. Preference-informed fairness. *arXiv preprint arXiv:1904.01793* (2019).
- [20] Matthäus Kleindessner, Pranjal Awasthi, and Jamie Morgenstern. 2020. A Notion of Individual Fairness for Clustering. *arXiv preprint arXiv:2006.04960* (2020).
- [21] Matt J Kusner, Joshua R Loftus, Chris Russell, and Ricardo Silva. 2017. Counterfactual fairness. *arXiv preprint arXiv:1703.06856* (2017).
- [22] Sepideh Mahabadi and Ali Vakilian. 2020. Individual fairness for k-clustering. In *International Conference on Machine Learning*. PMLR, 6586–6596.
- [23] Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena, Kristina Lerman, and Aram Galstyan. 2021. A survey on bias and fairness in machine learning. *ACM Computing Surveys (CSUR)* 54, 6 (2021), 1–35.
- [24] Alexandra Olteanu, Carlos Castillo, Fernando Diaz, and Emre Kiciman. 2019. Social data: Biases, methodological pitfalls, and ethical boundaries. *Frontiers in Big Data* 2 (2019), 13.
- [25] Akshat Pandey and Aylin Caliskan. 2021. Disparate Impact of Artificial Intelligence Bias in Ridehailing Economy’s Price Discrimination Algorithms. In *AIES ’21: AAAI/ACM Conference on AI, Ethics, and Society, Virtual Event, USA, May 19-21, 2021*, Marion Fourcade, Benjamin Kuipers, Seth Lazar, and Deirdre K. Mulligan (Eds.). ACM, 822–833.
- [26] Evangelia Pitoura, Kostas Stefanidis, and Georgia Koutrika. 2021. Fairness in Rankings and Recommendations: An Overview. *arXiv preprint arXiv:2104.05994* (2021).
- [27] Christopher Riederer and Augustin Chaintreau. 2017. The Price of Fairness in Location Based Advertising. (2017).
- [28] Philip B Stark and Robert L Parker. 1995. Bounded-variable least-squares: an algorithm and applications. *Computational Statistics* 10 (1995), 129–129.
- [29] Harini Suresh and John V Guttag. 2019. A framework for understanding unintended consequences of machine learning. *arXiv preprint arXiv:1901.10002* 2 (2019).